

# mogsa: gene set analysis on multiple omics data

*Chen Meng*

Modified: January 19, 2016. Compiled: October 7, 2025.

## Contents

|     |                                              |    |
|-----|----------------------------------------------|----|
| 1   | MOGSA overview. . . . .                      | 1  |
| 2   | Run mogsa . . . . .                          | 2  |
| 2.1 | Quick start . . . . .                        | 2  |
| 2.2 | Result analysis and interpretation . . . . . | 3  |
| 2.3 | Plot gene sets in projected space . . . . .  | 8  |
| 2.4 | Perform MOGSA in two steps . . . . .         | 8  |
| 3   | Preparation of gene set data . . . . .       | 10 |
| 4   | Session info . . . . .                       | 11 |

## 1 MOGSA overview

---

Modern "omics" technologies enable quantitative monitoring of the abundance of various biological molecules in a high-throughput manner, accumulating an unprecedented amount of quantitative information on a genomic scale. Gene set analysis is a particularly useful method in high throughput data analysis since it can summarize single gene level information into the biological informative gene set levels. The *mogsa* provide a method doing gene set analysis based on multiple omics data that describes the same set of observations/samples.

MOGSA algorithm consists of three steps. In the first step, multiple omics data are integrated using multi-table multivariate analysis, such as multiple factorial analysis (MFA) [1]. MFA projects the observations and variables (genes) from each dataset onto a lower dimensional space, resulting in sample scores (or PCs) and variables loadings respectively. Next, gene set annotations are projected as additional information onto the same space, generating a set of scores for each gene set across samples [2]. In the final step, MOGSA generates a gene set score (GSS) matrix by reconstructing the sample scores and gene set scores. A high GSS indicates that gene set and the variables in that gene set have measurement in one or more dataset that explain a large proportion of the correlated information across data tables. Variables (genes) unique to individual datasets or common among matrices may contribute to a high GSS. For example, in a gene set, a few genes may have high levels of gene expression, others may have increased protein levels and a few may have amplifications in copy number.

In this document, we show with an example how to use MOGSA to integrate and annotate multiple omics data.

## 2 Run mogsa

### 2.1 Quick start

In this working example, we will analyze the NCI-60 transcriptomic data from 4 different microarray platforms. The goal is to explore which functions (gene sets) are associated with (high or low expressed) which type of tumor. First, load the library and data

```
# loading gene expression data and supplementary data
library(mogsa)
library(ggplots) # used for visualizing heatmap
# loading gene expression data and supplementary data
data(NCI60_4array_supdata)
data(NCI60_4arrays)
```

`NCI60_4arrays` is a *list of data.frame*. The *list* consists of microarray data for NCI-60 cell lines from different platforms. In each of the *data.frame*, columns are the 60 cell lines and rows are genes. The data was downloaded from [3], but only a small subset of genes were selected. Therefore, the result in this vignette is not intended for biological interpretation.

`NCI60_4array_supdata` is a *list of matrix*, representing gene set annotation data. For each of the microarray data, there is a corresponding annotation matrix. In the annotation data, the rows are genes (in the same order as their original dataset) and columns are gene sets. An annotation matrix is a binary matrix, where 1 indicates a gene is present in a gene set and 0 otherwise. See the "Preparation of gene set data" section about how to create the gene set annotation matrices as required by `mogsa`. To have an overview of the two datasets:

```
sapply(NCI60_4arrays, dim) # check dimensions of expression data
##      agilent hgu133 hgu133p2 hgu95
## [1,]    300    298    268    288
## [2,]     60     60     60     60

sapply(NCI60_4array_supdata, dim) # check dimensions of supplementary data
##      agilent hgu133 hgu133p2 hgu95
## [1,]    300    298    268    288
## [2,]    150    150    150    150

# check if the gene expression data and annotation data are matched in the same order
identical(names(NCI60_4arrays), names(NCI60_4array_supdata))
## [1] TRUE

head(rownames(NCI60_4arrays$agilent)) # the type of gene IDs
## [1] "ST8SIA1" "YWHAQ" "EPHA4" "GTPBP5" "PVR" "ATP6V1H"
```

Also, we need to confirm the columns between the expression data and annotation data are mapped in the same order. To verify this, we do

```
dataColNames <- lapply(NCI60_4arrays, colnames)
supColNames <- lapply(NCI60_4arrays, colnames)
identical(dataColNames, supColNames)

## [1] TRUE
```

Before applying MOGSA, we first define a factor describing the tissue of origin of cell lines and color code, which will be used later.

```
# define cancer type
cancerType <- as.factor(substr(colnames(NCI60_4arrays$agilent), 1, 2))
# define color code to distinguish cancer types
colcode <- cancerType
levels(colcode) <- c("black", "red", "green", "blue",
                    "cyan", "brown", "pink", "gray", "orange")
colcode <- as.character(colcode)
```

Then, we call the function `mogsa` to run MOGSA:

```
mgsa1 <- mogsa(x = NCI60_4arrays, sup=NCI60_4array_supdata, nf=3,
               proc.row = "center_ssqr", w.data = "inertia", statis = TRUE)
```

In this function, the input argument `proc.row` stands for the preprocessing of rows and argument `w.data` indicates the weight of datasets. The last argument `statis` is about which multiple table analysis method should be used. Two multivariate methods are available at present, one is "STATIS" (`statis=TRUE`) [4], the other one is multiple factorial analysis (MFA; `statis=FALSE`, the default setting) [1].

In this analysis, we arbitrarily selected top three PCs (`nf=3`). But in practice, the number of PCs need to be determined before running the MOGSA. Therefore, it is also possible to run the multivariate analysis and projecting annotation data separately. After running the multivariate analysis, a scree plot of eigenvalues for each PC could be used to determine the proper number of PCs to be included in the annotation projection step (See the "Perform MOGSA in two steps" section).

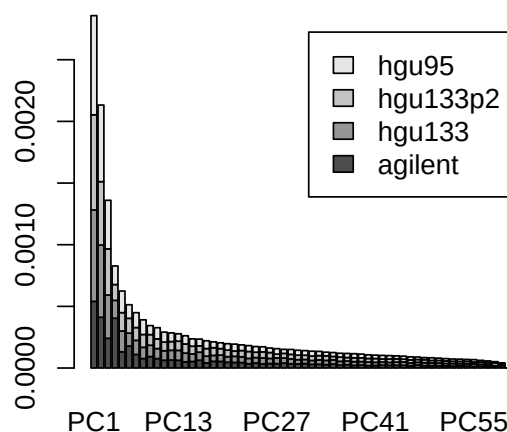
## 2.2 Result analysis and interpretation

The function `mogsa` returns an object of class `mgsa`. This information could be extracted with function `getmgsa`. First, we want to know the variance explained by each PC on different datasets (figure 1).

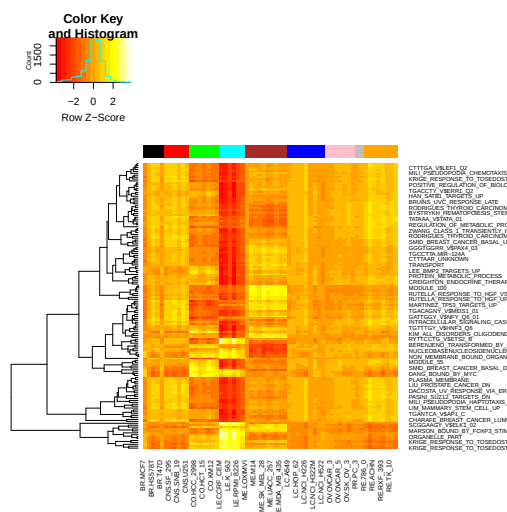
```
eigs <- getmgsa(mgsa1, "partial.eig") # get partial "eigenvalue" for separate data
barplot(as.matrix(eigs), legend.text = rownames(eigs))
```

The main result returned by `mogsa` is the gene set score (GSS) matrix. The value in the matrix indicates the overall active level of a gene set in a sample. The matrix could be extracted and visualized by

```
# get the score matrix
scores <- getmgsa(mgsa1, "score")
heatmap.2(scores, trace = "n", scale = "r", Colv = NULL, dendrogram = "row",
          margins = c(6, 10), ColSideColors=colcode)
```



**Figure 1:** The variance of each principal components (PC), the contributions of different data are distinguished by different colors



**Figure 2:** heatmap showing the gene set score (GSS) matrix

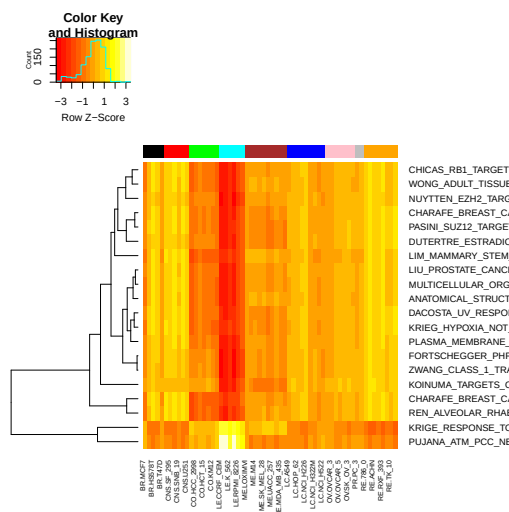
Figure 2 shows the gene set score matrix returned by **mogsa**. The rows of the matrix are all the gene sets used to annotate the data. But we are mostly interested in the gene sets with large number of significant gene sets, because these gene sets describe the difference across cell lines. The corresponding p-value for each gene set score could be extracted by **getmgsa**. Then, the most significant gene sets could be defined as gene sets that contain highest number of significantly p-values. For example, if we want to select the top 20 most significant gene sets and plot them in heatmap, we do:

```
p.mat <- getmgsa(mgsa1, "p.val") # get p value matrix
# select gene sets with most significant GSS scores.
top.gs <- sort(rowSums(p.mat < 0.01), decreasing = TRUE)[1:20]
top.gs.name <- names(top.gs)
top.gs.name

## [1] "PASINI_SUZ12_TARGETS_DN"
## [2] "CHARAFE_BREAST_CANCER_LUMINAL_VS_BASAL_DN"
## [3] "CHARAFE_BREAST_CANCER_LUMINAL_VS_MESENCHYMAL_DN"
```

```
## [4] "KOINUMA_TARGETS_OF_SMAD2_OR_SMAD3"
## [5] "DUTERTRE_ESTRADIOL_RESPONSE_24HR_DN"
## [6] "REN_ALVEOLAR_RHABDOMYOSARCOMA_DN"
## [7] "LIM_MAMMARY_STEM_CELL_UP"
## [8] "LIU_PROSTATE_CANCER_DN"
## [9] "CHICAS_RB1_TARGETS_CONFLUENT"
## [10] "NUYTEN_EZH2_TARGETS_UP"
## [11] "DACOSTA_UV_RESPONSE_VIA_ERCC3_DN"
## [12] "PUJANA_ATM_PCC_NETWORK"
## [13] "KRIGE_RESPONSE_TO_T0SED0STAT_24HR_DN"
## [14] "WONG_ADULT_TISSUE_STEM_MODULE"
## [15] "KRIEG_HYPOXIA_NOT_VIA_KDM3A"
## [16] "MULTICELLULAR_ORGANISMAL_DEVELOPMENT"
## [17] "ANATOMICAL_STRUCTURE_DEVELOPMENT"
## [18] "FORTSCHEGGER_PHF8_TARGETS_DN"
## [19] "ZWANG_CLASS_1_TRANSIENTLY_INDUCED_BY_EGF"
## [20] "PLASMA_MEMBRANE_PART"
```

```
heatmap.2(scores[top.gs.name, ], trace = "n", scale = "r", Colv = NULL, dendrogram = "row",
  margins = c(6, 10), ColSideColors=colcode)
```



**Figure 3:** heatmap showing the gene set score (GSS) matrix for top 20 significant gene sets

The result is shown in figure 3. We can see that these gene sets reflect the difference between leukemia and other tumors.

So far, we already had an integrative overview of gene sets active levels over the 60 cell lines. It is also interesting to look into more detailed information for a specific gene set. For example, which dataset(s) contribute most to the high or low gene set score of a gene set? And which genes are most important in defining the gene set score for a gene set? The former question could be answered by the gene set score decomposition; the later question could be solve by the gene influential score. These analysis can be done with [decompose.gs.group](#) and [GIS](#).

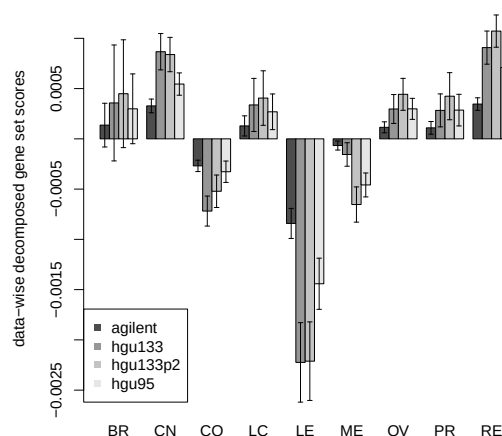
In the first example, we explore the gene set that have most significant gene set scores. The gene set is

```
# gene set score decomposition
# we explore two gene sets, the first one
gs1 <- top.gs.name[1] # select the most significant gene set
gs1

## [1] "PASINI_SUZ12_TARGETS_DN"
```

The data-wise decomposition of this gene set over cancer types is

```
# decompose the gene set score over datasets
decompose.gs.group(mgsa1, gs1, group = cancerType)
```



**Figure 4:** gene set score (GSS) decomposition. The GSS decomposition are grouped according to the tissue of origin of cell lines. The vertical bar showing the 95% of confidence interval of the means.

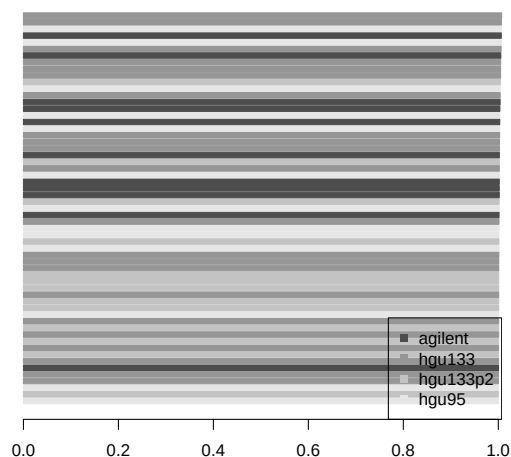
Figure 4 shows leukemia cell lines have lowest GSS on this gene set. The contribution to the overall gene set score by each dataset are separated in this plot. In general, there is a good concordance between different datasets. But HGU133 platform contribute most and Agilent platform contributed least comparing with other datasets, represented as the longest or shortest bars.

Next, in order to know the most influential genes in this gene set. We call the function **GIS**:

```
gis1 <- GIS(mgsa1, gs1, barcol = gray.colors(4)) # gene influential score
```

```
head(gis1) # print top 6 influencers
```

```
## feature    GIS    data
## 1  LIMD2 1.007091 hgu133
## 2  ZNF266 1.006706 hgu133
## 3  LIMD2 1.006476 hgu95
## 4   GNG2 1.006327 agilent
## 5   SP5 1.006035 hgu95
## 6   SP5 1.005954 hgu133
```

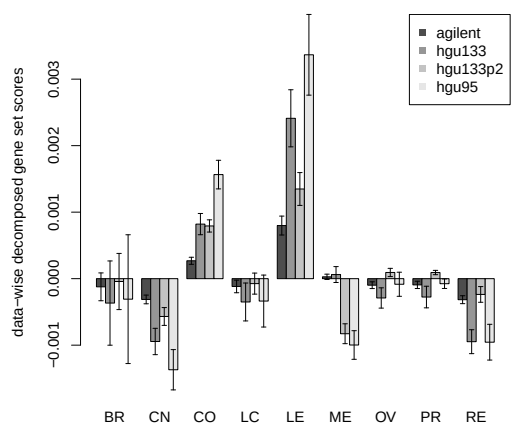


**Figure 5:** The gene influential score (GIS) plot. the GIS are represented as bars and the original data where the gene is from is distinguished by different colors.

In figure 5, the bars represent the gene influential scores for genes. Genes from different platforms are shown in different colors. The expression of genes with high positive GIS more likely to have a good positive correlation with the gene set score. In this example, the most important genes in the gene set "PASIN SUZ12 TARGETS DN" are TNFRSF12A (identified in two different platforms), CD151, ITGB1, etc.

In the next example, we use the same methods to explore the "PUJANA ATM PCC NETWORK" gene set.

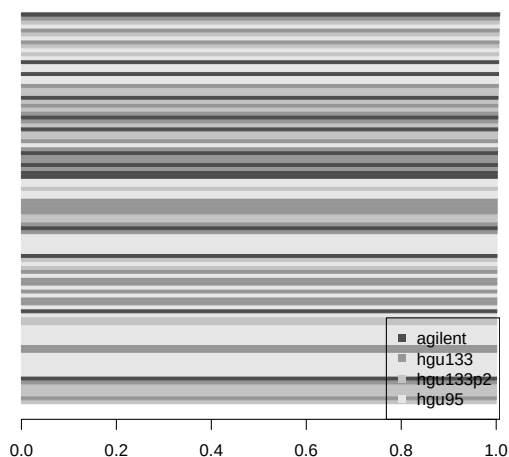
```
# the section gene set
gs2 <- "PUJANA_ATM_PCC_NETWORK"
decompose.gs.group(mgsa1, gs2, group = cancerType, x.legend = "topright")
```



**Figure 6:** Data-wise decomposed GSS for gene set 'PUJANA ATM PCC NETWORK'

```
gis2 <- GIS(mgsa1, "PUJANA_ATM_PCC_NETWORK", topN = 6, barcol = gray.colors(4))
```

```
gis2
## feature      GIS      data
## 1  PCBP4 1.007281 agilent
```



**Figure 7:** GIS plot for gene set 'PUJANA ATM PCC NETWORK'

```
## 2    LIF 1.006737  hgu133
## 3    DKK3 1.006393 hgu133p2
## 4    ROB01 1.006231  hgu95
## 5    GPD2 1.006213  hgu133
## 6    KCNMA1 1.006116 hgu133p2
```

Figure 6 shows that the leukemia cell lines have highest GSSs for this gene set. And the HGU133 and HGU95 platform have relative high contribution to the overall gene set score. The GIS analysis (figure 7) indicates the PIK4CG and GMFG are the most important genes in this gene set.

## 2.3 Plot gene sets in projected space

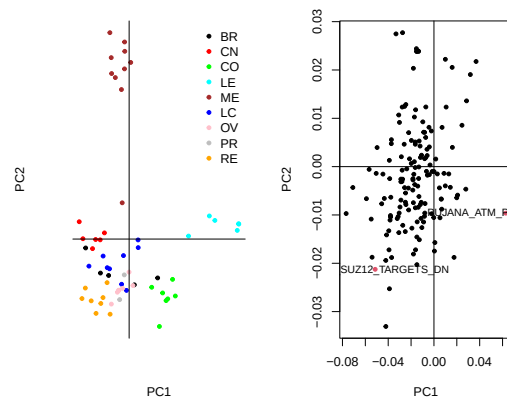
We can also see how the gene set are presented in the lower dimension space. Here we show the projection of gene set annotations on first two dimensions. Then, the label the two gene sets we analyzed before.

```
fs <- getmgsa(mgsa1, "fac.scr") # extract the factor scores for cell lines (cell line space)
layout(matrix(1:2, 1, 2))
plot(fs[, 1:2], pch=20, col=colcode, axes = FALSE)
abline(v=0, h=0)
legend("topright", col=unique(colcode), pch=20, legend=unique(cancerType), bty = "n")
plotGS(mgsa1, label.cex = 0.8, center.only = TRUE, topN = 0, label = c(gs1, gs2))
```

## 2.4 Perform MOGSA in two steps

**mogsa** perform MOGSA in one step. But in practice, one need to determine how many PCs should be retained in the step of reconstructing gene set score matrix. A scree plot of the eigenvalues, which result from the multivariate analysis, could be used for this purpose. Therefore, we can perform the multivariate data analysis and gene set annotation projection in two steps. To do the multivariate analysis, we call the **moa**:



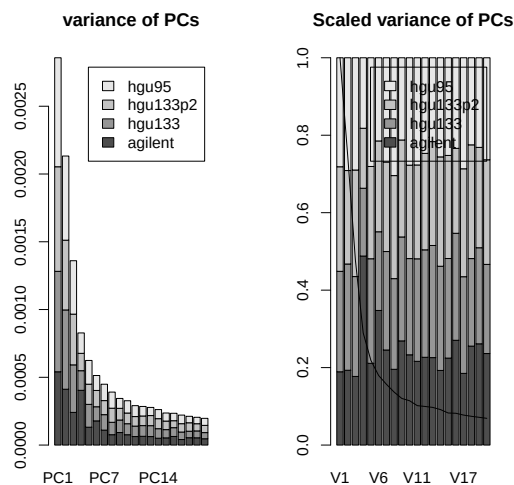


**Figure 8:** cell line and gene sets projected on the PC1 and PC2

```
# perform multivariate analysis
ana <- moa(NCI60_4arrays, proc.row = "center_ssq1", w.data = "inertia", statis = TRUE)
slot(ana, "partial.eig")[, 1:6] # extract the eigenvalue

##          PC1          PC2          PC3          PC4          PC5
## agilent  0.0005406833 0.0004119778 0.0002410063 0.0004038087 0.0001317894
## hgu133   0.0007410830 0.0005850680 0.0003507538 0.0001448788 0.0001685482
## hgu133p2 0.0007716595 0.0005146566 0.0003742008 0.0001281515 0.0001487516
## hgu95    0.0008042677 0.0006210049 0.0003942394 0.0001506287 0.0001752495
##          PC6
## agilent  0.0001783712
## hgu133   0.0001042850
## hgu133p2 0.0001203610
## hgu95    0.0001102364

# show the eigenvalues in scree plot:
layout(matrix(1:2, 1, 2))
plot(ana, value="eig", type = 2, n=20, main="variance of PCs") # use '? "moa-class"' to check the help manu
plot(ana, value="tau", type = 2, n=20, main="Scaled variance of PCs")
```



**Figure 9:** cell line and gene sets projected on the PC1 and PC2

The multivariate analysis (`moa`) returns an object of class `moa-class`. The scree plot shows the top 3 PC is the most significant since they explain much more variance than others. Several other methods, such as the informal "elbow test" or more formal test could be used to determine the number of retained PCs [5]. In order to be consistent with previous example, we use top 3 PCs in the analysis:

```

mogsa2 <- mogsa(x = ana, sup=NCI60_4array_supdata, nf=3)

## x is an object of "moa", statis is not used

identical(mogsa1, mogsa2) # check if the two methods give the same results

## [1] FALSE

```

### 3 Preparation of gene set data

Package `GSEABase` provides several methods to create a gene set list [6]. In `mogsa` there are two methods to create gene set list. The first one is generating gene set list from package `graphite` [7] using function `prepGraphite`.

```

library(graphite)
keggdb <- prepGraphite(db = pathways("hsapiens", "kegg")[1:50], id = "symbol")

## converting identifiers!
## converting identifiers done!

keggdb <- lapply(keggdb, function(x) sub("SYMBOL:", "", x))
keggdb[1:2]

## $`Glycolysis / Gluconeogenesis`
## [1] "AKR1A1" "ADH1A" "ADH1B" "ADH1C" "ADH4" "ADH5" "ADH6"
## [8] "GALM" "ADH7" "LDHAL6A" "DLAT" "DLD" "EN01" "EN02"
## [15] "EN03" "ALDH2" "ALDH3A1" "ALDH1B1" "FBP1" "ALDH3B1" "ALDH3B2"
## [22] "ALDH9A1" "ALDH3A2" "ALDOA" "ALDOB" "ALDOC" "G6PC1" "GAPDH"
## [29] "GAPDHS" "GCK" "GPI" "HK1" "HK2" "HK3" "EN04"
## [36] "LDHA" "LDHB" "LDHC" "PGAM4" "ALDH7A1" "PCK1" "PCK2"
## [43] "PDHA1" "PDHA2" "PDHB" "PFKL" "PFKM" "PFKP" "PGAM1"
## [50] "PGAM2" "PGK1" "PGK2" "PGM1" "PKLR" "PKM" "PGM2"
## [57] "ACSS2" "G6PC2" "BPGM" "TPI1" "HKDC1" "ADPGK" "ACSS1"
## [64] "FBP2" "LDHAL6B" "G6PC3" "MINPP1"
##
## $`Citrate cycle (TCA cycle)`
## [1] "CS" "DLAT" "DLD" "DLST" "FH" "IDH1" "IDH2" "IDH3A"
## [9] "IDH3B" "IDH3G" "MDH1" "MDH2" "ACLY" "AC01" "OGDH" "AC02"
## [17] "PC" "PDHA1" "PDHA2" "PDHB" "OGDHL" "SDHA" "SDHB" "SDHC"
## [25] "SDHD" "SUCLG2" "SUCLG1" "SUCLA2" "PCK1" "PCK2"

```

The second method is to create a gene set list from "gmt" files, which could be downloaded from MSigDB [8] after obtaining a proper license. In our working example, we will work on a toy example from this database containing only three datasets.

```

dir <- system.file(package = "mogsa")
preGS <- prepMSigDB(file=paste(dir, "/extdata/example_msigdb_data.gmt.gz", sep = ""))

```

In order to use the gene set information in **mogsa**, we have to convert the list of gene sets to a list of annotation matrix. This can be done with **prepSupMoa**. This function requires two obligatory inputs, first is the multiple omics datasets and the second input could be a gene set list, *GeneSet* or *GeneSetCollection*. The output of **prepSupMoa** could be directly passed into the **mogsa**.

```
# the prepare
sup_data1 <- prepSupMoa(NCI60_4arrays, geneSets=keggdb, minMatch = 1)
mgsa3 <- mogsa(x = NCI60_4arrays, sup=sup_data1, nf=3,
               proc.row = "center_ssq1", w.data = "inertia", statis = TRUE)
```

## 4 Session info

```
toLatex(sessionInfo())
```

- R version 4.5.1 Patched (2025-08-23 r88802), x86\_64-pc-linux-gnu
- Locale: LC\_CTYPE=en\_US.UTF-8, LC\_NUMERIC=C, LC\_TIME=en\_GB, LC\_COLLATE=C, LC\_MONETARY=en\_US.UTF-8, LC\_MESSAGES=en\_US.UTF-8, LC\_PAPER=en\_US.UTF-8, LC\_NAME=C, LC\_ADDRESS=C, LC\_TELEPHONE=C, LC\_MEASUREMENT=en\_US.UTF-8, LC\_IDENTIFICATION=C
- Time zone: America/New\_York
- TZcode source: system (glibc)
- Running under: Ubuntu 24.04.3 LTS
- Matrix products: default
- BLAS: /home/biocbuild/bbs-3.22-bioc/R/lib/libRblas.so
- LAPACK: /usr/lib/x86\_64-linux-gnu/lapack/liblapack.so.3.12.0
- Base packages: base, datasets, grDevices, graphics, methods, stats, utils
- Other packages: knitr 1.50, mogsa 1.43.0
- Loaded via a namespace (and not attached): AnnotationDbi 1.71.1, Biobase 2.69.1, BiocGenerics 0.55.1, BiocManager 1.30.26, BiocStyle 2.37.1, Biostrings 2.77.2, DBI 1.2.3, GSEABase 1.71.1, IRanges 2.43.5, KEGGREST 1.49.1, KernSmooth 2.23-26, Matrix 1.7-4, MatrixGenerics 1.21.0, R6 2.6.1, RSQLite 2.4.3, S4Vectors 0.47.4, Seqinfo 0.99.2, XML 3.99-0.19, XVector 0.49.1, annotate 1.87.0, bit 4.6.0, bit64 4.6.0-1, bitops 1.0-9, blob 1.2.4, caTools 1.18.3, cachem 1.1.0, cli 3.6.5, cluster 2.1.8.1, codetools 0.2-20, compiler 4.5.1, corpcor 1.6.10, crayon 1.5.3, digest 0.6.37, evaluate 1.0.5, fastmap 1.2.0, genefilter 1.91.0, generics 0.1.4, gplots 3.2.0, graph 1.87.0, graphite 1.55.0, grid 4.5.1, gtools 3.9.5, highr 0.11, htmltools 0.5.8.1, httr 1.4.7, lattice 0.22-7, matrixStats 1.5.0, memoise 2.0.1, parallel 4.5.1, png 0.1-8, rappdirs 0.3.3, rlang 1.1.6, rmarkdown 2.30, splines 4.5.1, stats4 4.5.1, survival 3.8-3, svd 0.5.8, tinytex 0.57, tools 4.5.1, vctrs 0.6.5, xfun 0.53, xtable 1.8-4, yaml 2.3.10

## References

- [1] Herve Abdi, Lynne J. Williams, and Dominique Valentin. Multiple factor analysis: principal component analysis for multitable and multiblock data sets. *Wiley Interdisciplinary Reviews: Computational Statistics*, 5:149–179, 2013.
- [2] M. de Tayrac, S. Le, M. Aubry, J. Mosser, and F. Husson. Simultaneous analysis of distinct omics data sets with integration of biological knowledge: Multiple factor analysis approach. *BMC Genomics*, 10:32, 2009.
- [3] Reinhold WC, Sunshine M, Liu H, Varma S, Kohn KW, Morris J, Doroshow J, and Pommier Y. Cellminer: A web-based suite of genomic and pharmacologic tools to explore transcript and drug patterns in the nci-60 cell line set. *Cancer Research*, 72(14):3499–511, 2012.
- [4] Herve Abdi, Lynne J. Williams, Dominique Valentin, and Mohammed Bennani-Dosse. Statis and distatis: optimum multitable principal component analysis and three way metric multidimensional scaling. *Wiley Interdisciplinary Reviews: Computational Statistics*, 4:124–167, 2012.
- [5] Herve Abdi and Lynne J. Williams. Principal component analysis. *Wiley Interdisciplinary Reviews: Computational Statistics*, 2:433–459, 2010.
- [6] Morgan M, Falcon S, and Gentleman R. Gseabase: Gene set enrichment data structures and methods. *R package version 1.28.0*.
- [7] Gabriele Sales<sup>1</sup>, Enrica Calura<sup>1</sup>, Duccio Cavalieri, and Chiara Romualdi<sup>1</sup>. graphite - a bioconductor package to convert pathway topology to gene network. *BMC bioinformatics*, 13:20, 2012.
- [8] Aravind Subramanian, Pablo Tamayo, Vamsi K. Mootha, Sayan Mukherjee, Benjamin L. Ebert, Michael A. Gillette, Amanda Paulovich, Scott L. Pomeroy, Todd R. Golub, Eric S. Lander, and Jill P. Mesirov. Gene set enrichment analysis: A knowledge-based approach for interpreting genome-wide expression profiles. *Proceedings of the National Academy of Sciences*, 102:15545–15550, 2005.